



Leading the People Side of AI

How leaders, managers, and teams can help AI create real value

Acknowledgement: As the dawning reality of AI hits, we all have strong feelings and legitimate concerns. This paper is not advocating for AI use. Rather it acknowledges the need for organizations to recover their AI investments. And it steers that momentum in ways that deliver durable value—both for organizations and humans. This paper shares the lessons FourSight has learned, supporting innovation over the last 20 years. AI is simply the latest one. We know that people can innovate when they know how. Here's what we know about how to harness AI to create value.

EXECUTIVE SUMMARY

AI doesn't create value.

People using AI effectively do.

Specifically, people who share a unique skillset: They know how to recognize worthwhile problems, push beyond obvious ideas, visualize and prototype solutions, and act on the best ones.

Those aren't technical skills. They are thinking skills.

Without strong thinking skills, people using AI tend to...

- Focus on the wrong problems.
- Scale the wrong solutions.

Organizations can create value with AI when 1) people share a common language for the thinking process that leads to value creation (aka innovation), and 2) leaders foster a healthy dynamic between top-down direction and bottom-up discovery.

Managing the dynamic tension between organizational control and individual discovery is the best way to avoid swinging between the disappointing extremes of individual AI experiments that don't scale and centralized agents that miss the mark.

This paper describes three positions can connect organizational direction with individual and team learning: Leaders, Managers, and Individuals/Teams.

- **Leaders** identify strategic goals, set boundaries, resource learning.
- **Individuals /Teams** discover problems, prototype and validate solutions.
- **Managers** translate in both directions without becoming the sole gatekeepers

AI literacy is necessary, but it is not sufficient. This paper focuses on elevating the durable skills of human problem solving so the best AI solutions are discovered, put forward, improved and scaled.

That means building skills of metacognition and creative problem solving. It means giving Individuals and Teams structured opportunities to discover problems, prototype solutions, and validate what helps. It means preparing Managers to frame, translate, amplify, and reduce noise. It means asking Leaders to monitor the ecosystem, focus attention, resource learning, and decide broader scope.

This approach does not replace AI governance, change management, product or service ownership, or technical operations.

Instead, it engages people in how to put AI to work for them.

A note on how we wrote this whitepaper:

This paper is the culmination of interviews, research, and experience supporting innovation and AI governance. We used AI to summarize and format the thinking, which is entirely our own. Our process began with research, interviews, debate, creative problem solving, and whiteboarding. Finally, we recorded a deliberate conversation, articulating all our ideas and methods and fed the output to AI to wordsmith and format. We hope this document provides value to you, and we look forward to your feedback.

1. AI adoption is not just a technology rollout

Most workplace tools arrive with a job description. A spreadsheet calculates. A CRM stores customer information. A project management system tracks work. AI arrives with a menu of possible jobs. It can draft, code, summarize, analyze, compare, simulate, translate, question, and prototype.

That makes AI powerful. It also makes AI confusing.

A vague instruction to “use AI” produces activity but not necessarily progress. Some people find useful shortcuts. Others produce polished demonstrations that solve low-value problems. Some create tools that work for one person but are risky, redundant, or disruptive to adjacent work. Usage increases without a clear account of what became better.

Centralized AI programs create the opposite failure mode. When strategy lives only with executives, IT, or a central innovation group, solutions can be designed too far from the work. The people doing the job often know where the friction lives: the approval that delays everything, the report rebuilt by hand, the customer handoff that repeatedly fails, or the rule that is technically followed but operationally painful.

Discovery without integration creates fragmentation. Control without discovery misses the work.

A people-centered approach begins with a different assumption: AI value emerges when organizational direction meets local knowledge, and when learning from local work can influence organizational decisions. When organizations manage the ongoing dynamic between “organizational direction” and “individual discovery.”

Leaders provide direction. Individuals and teams provide context, prototypes, and evidence. Managers help meaning move between them. Relevant domain expertise, authority, and resources join as the scope requires. AI helps people explore and test new ways of working; people remain responsible for deciding what the work means and what should happen next.

2. The missing capability: metacognition

AI has made thinking more visible.

When a person works with AI, they reveal how they understand the problem. A weak prompt may reflect an unclear goal, missing context, an untested assumption, or a narrow definition of success. AI can respond quickly, but speed is not the same as value. A plausible answer can make weak thinking look finished before anyone has tested whether it is useful.

This is why metacognition matters.

Metacognition is the ability to think about one’s thinking. In an AI-enabled workplace, it helps people pause and ask:

- What problem are we really trying to solve?
- What do we know, and what do we still need to understand?
- Whose experience or expertise is missing?
- What kind of thinking does the work need now?
- Should we clarify the challenge, generate ideas, develop a solution, implement it, or learn from what happened?
- What evidence would make the result credible?

These questions turn AI from a shortcut into a resource for problem solving.

For individuals, metacognition improves framing, prompting, evaluation, and learning. For teams, it creates a language for recognizing different thinking preferences and choosing the kind of thinking the work needs next. For organizations, it makes experimentation easier to understand, compare, and connect.

AI can also change the unit of contribution. A person who once could only describe an idea may now bring a rough workflow, checklist, interface, analysis, agent, or internal tool. The prototype may be fragile or incomplete. That is acceptable when its purpose is learning. A tangible prototype changes the conversation from “Do we like this idea?” to “What happened when someone tried it?”

This greater reach does not create universal judgment. People may prototype outside their expertise and explore interdisciplinary possibilities before the organization commits scarce resources. But a useful prototype may still be insecure, inaccessible, noncompliant, off-brand, technically fragile, or unsuitable for wider use. Better AI-enabled work depends on knowing both what one can explore and where other expertise is needed.

AI does not replace human thinking. It raises the value of understanding how thinking works.

3. The polarity: direction and discovery

Every organization using AI must manage a persistent polarity.

On one side is organizational direction and integration. The organization needs strategic focus, shared infrastructure, security, governance, quality, consistency, ownership, and a way to allocate limited resources. It cannot treat every prototype as an official solution.

On the other side is individual and team discovery and learning. People closest to the work need room to explore problems, test possibilities, and improve the work they know. They see the exceptions, workarounds, delays, customer needs, and weak links that formal process maps often miss.

If the organization overemphasizes local discovery, AI efforts can fragment. If it overemphasizes centralized direction, AI efforts can become detached from reality. The answer is not a permanent compromise in the middle. It is a continuing exchange between the two sides.

Strategic goals, risks, and boundaries move toward the work. Prototypes, evidence, controls, and learning move toward wider organizational awareness. Managers add translation in both directions, but the exchange is not limited to a single chain of command. Evidence should remain discoverable to peers, affected people, relevant domain experts, and the decision forums appropriate to its scope.

The same exchange can operate across a whole organization, inside a domain, within a team, or in a temporary initiative. A person may be a Manager in one exchange, a Leader within a domain, and an Individual or Team member doing the work in another. These are positions in the work, not three fixed classes of employee.

Work may begin with a strategic concern, a local problem, a useful prototype, an operational bottleneck, or a governance failure. Learning can send it backward, sideways, into a wider scope, or to retirement. A healthy model must allow that movement.

4. A people-centered AI adoption model

A useful adoption model clarifies how direction and learning move without pretending that every organization has the same structure.

The model has three actor positions: Leaders, Managers, and Individuals/Teams. It also recognizes many domains of expertise. Expertise, responsibility, and authority are related, but they are not the same thing.

Leaders: Set direction and decide scope

Leaders monitor the ecosystem around the organization—customers, competitors, technology, regulation, workforce, and other forces—and identify the strategic goals and risks that deserve attention. They make those priorities legible, establish useful boundaries, resource learning, and decide which evidence-backed opportunities warrant broader investment or use.

Their job is not to demand generic AI usage, prescribe every local challenge, or write every operational rule. A Leader may recognize data leakage as a strategic risk without knowing the exact control needed in every domain. Leaders set direction and call attention to risk; people closer to the work translate that direction into challenges, requirements, controls, and solutions.

Leaders should be cautious with usage-based vanity metrics. Usage may diagnose access, participation, or abandonment, but it is not evidence of value by itself.

Managers: Frame, amplify, and reduce noise

Managers connect organizational direction with local work.

Moving toward the team, they translate broad goals and concerns into challenges people can understand and act on. Moving toward broader organizational awareness, they help tailor validation, frame evidence for its audience, amplify strong signal, and reduce noise.

Managers strengthen evidence; they do not own it. Manager enthusiasm can accelerate an initiative, but one Manager should not be the only path by which a validated prototype or its learning remains visible. Nor should every local experiment require Manager permission. The need for consultation or decision depends on the scope, risk, affected work, and resources involved.

Managers can ask what would make a test convincing, who should participate, how the work connects to priorities, and how the evidence should reach its next audience. They are framers and amplifiers, not kingmakers.

Individuals and Teams: Discover, prototype, and validate

Individuals and Teams know the details of work: the handoffs, exceptions, workarounds, unwritten rules, recurring frustrations, and hidden opportunities.

They discover problems, prototype solutions, validate whether those solutions help in context, improve the work, and communicate what they learned. They own the experiment and the evidence it produces. The evidence may be positive, negative, mixed, or inconclusive. Negative learning is still valuable when it prevents poor investment or reveals a better problem.

A prototype is built to learn. It may remain useful to one person or one team. It may reveal an opportunity worth broader attention. It may fail. The model does not assume that every prototype should become a product.

Everyone is a domain expert in something

Domain Expert is not a separate role. Leaders, Managers, and Individuals/Teams all possess expertise in particular contexts. A facilities employee, marketer, engineer, customer-service representative, manager, or executive knows something about work that others may not.

Expertise is relative to the work. A systems engineer may verify systems logic but not marketing practice. A marketing expert may define brand requirements but not production-software reliability. An AI-governance expert may evaluate a control's monitoring and failure modes but still need the affected domain to define what correct work means.

People may contribute domain expertise, build an MVP, improve a process, define a requirement, or verify a control from any of the three actor positions. That is why actor position should not be confused with profession or organizational level.

Scope determines the help required

If work stays within a team's expertise and authority, it can often be prototyped and validated locally. Manager consultation is useful when it improves framing, evidence, or connection, but it is not automatically a gate.

As work affects another domain, reaches more people, requires more coordination, crosses authority boundaries, or carries greater risk, broader collaboration becomes necessary. The organization should ask whether the people doing the work have the needed expertise, who else would be affected, who has decision authority and resources, and what evidence is proportionate to the risk.

Teams should not be expected to recognize every boundary unaided. Organizations should make prohibited conditions and routing signals visible and, where appropriate, enforce them through approved tools, data boundaries, and environments. The purpose is to help work move responsibly, not to make every concern a blocking activity.

5. From discovery to responsible use

The work is better understood as a set of recurring activities than as a mandatory sequence of phases. An organization may begin with strategic sensing, a local problem, a prototype, a failed control, or a need for wider adoption. Activities may overlap, repeat, or send the work in a new direction.

Strategic sensing and framing

Leaders monitor the environment and identify goals, risks, and questions that matter to the organization. They communicate meaningful intent, boundaries, and value criteria without prescribing a generic amount of AI use or every local solution.

The output is not an AI mandate. It is clearer organizational attention: what needs to improve, what concerns deserve investigation, where learning is needed, and what limits are already known.

Local framing and discovery

Managers help connect strategic concerns to local work, while Individuals and Teams can also begin with problems they encounter directly. Discovery includes people affected by the work, explores alternative problem frames, and considers non-AI responses before a prototype narrows the conversation.

The aim is not a long use-case inventory. It is a better understanding of a worthwhile challenge and the smallest useful way to learn about it.

Prototyping and validation

Individuals and Teams prototype solutions and test whether they help. They compare what happened with a baseline, expectation, or current way of working; invite relevant people to try the prototype; observe benefits, failures, shifted burdens, and unintended effects; and state what remains uncertain.

Validation asks whether a solution helps in the world. The Individual or Team owns the experiment and its evidence. As the affected scope grows, people whose work or interests may change should help shape the validation, and relevant owners should decide what acceptance means.

Evidence framing and portfolio learning

Managers help make evidence understandable and relevant without taking ownership of it. Leaders can then compare evidence-bearing opportunities rather than sort idea lists by enthusiasm.

Portfolio decisions should consider strategic fit, evidence quality, system effects, risk, total cost, and organizational capacity. Similar prototypes may be compared, combined, or used to reveal a broader systems issue. The most useful convergence clusters what the prototypes taught, not only what they were called.

MVP development and verification

A minimum viable product is not simply a polished prototype. It is a more disciplined solution built for broader responsibility by people whose combined expertise covers the relevant domains.

Validation and verification remain distinct. Validation asks whether the solution helps. Verification asks whether it meets the applicable requirements. Those requirements may concern technology, data, security, accessibility, law, policy, brand, operations, or the substance of another domain. A concept can appear useful before every requirement is verified, but broader or higher-risk use requires explicit expertise, verification evidence, limitations, and open risks.

Automated checks may perform verification only when the organization has evidence that the control works and Leaders authorize its use. People still define correctness, operating limits, and monitoring.

Domain governance

Domain governance turns recurring concerns into reusable capability. People with relevant expertise translate policy, law, professional standards, and organizational priorities into requirements, tests, templates, filters, escalation paths, approved components, and other controls.

These controls inject quality and compliance into future prototyping, building, validation, and verification. They can reduce repeated review and improve the capacity of people doing the work. Governance is therefore enabling infrastructure, not merely a checkpoint at the end and not a universal reason to stop discovery.

Wider pilot, ownership, and renewed learning

Verification does not prove that a solution will remain useful in a different context. When an MVP reaches more people, affects more work, or changes operating conditions, it must be validated again.

Before a pilot becomes normal work, name its owner, support, monitoring, incident response, review, and retirement path. A useful and verified improvement still needs a transition shaped by the people whose work it changes.

6. Training for AI-enabled problem solving

A one-size-fits-all AI workshop will not be enough. The durable capability is not prompt craft. It is the ability to understand a problem, make possibilities tangible, test them, recognize judgment boundaries, and involve the right people as scope changes.

A shared foundation: Creative problem solving and metacognition

Everyone benefits from a common language for the creative process. People should be able to recognize whether the work needs clarification, ideation, solution development, or implementation—and when learning requires a return to another kind of thinking.

Training should also help people understand their own thinking preferences and appreciate the preferences others bring. This makes AI use less about finding a perfect prompt and more about deliberately choosing the thinking the work needs.

Leadership learning: Strategic direction and portfolio learning

Leaders need practice in ecosystem sensing, strategic framing, risk communication, value criteria, resource allocation, and portfolio learning. They should learn to turn changes in customers, competition, technology, regulation, and workforce conditions into clear goals and concerns without prescribing local answers.

Leadership learning should address AI's cognitive and organizational effects, the difference between usage and value, the evidence needed for broader commitment, and the boundaries for local discovery. The aim is focused attention and responsible investment, not quotas.

Manager learning: Translation, validation consultation, and evidence

Managers need practice translating strategic direction into useful local challenges, discovering weak links, and helping teams frame work without taking it over. They should be able to consult on validation design, identify who should participate, recognize shifted burdens, and help teams communicate evidence to different audiences.

Managers also need to practice amplification and noise reduction: making learning visible, connecting similar efforts, and distinguishing evidence from novelty without becoming the sole approval path.

Individual and Team learning: Discovery, prototyping, and validation

Individuals and Teams need hands-on experience using AI in real work. They should learn the approved tools, data boundaries, and action limits; clarify challenges; generate and develop alternatives; build small prototypes; test usefulness; identify affected people and domains; recognize limitations and expertise gaps; and communicate evidence honestly.

A well-designed discovery sprint teaches the creative problem-solving process, helps participants see their thinking preferences, and asks them to work on meaningful challenges within organizational goals and rules. Participants explore where AI helps and where it does not. They leave with more than a list of ideas: they leave with tangible minimal prototypes, early observations, and a clear next learning step.

Domain-focused capability: Requirements, controls, and verification

Domain-focused learning is not reserved for a separate class of experts. It helps people apply the expertise they already hold to AI-enabled work. Depending on the domain, this may include process improvement, requirements definition, reusable controls, verification methods, MVP development, monitoring, incident response, and cross-domain collaboration.

People who build controls should learn how to make correct practice easier upstream, not only how to detect failure at the end. People who verify work should understand the evidence their domain requires and the limits of automated checks. People joining an MVP effort should understand how responsibility changes as a prototype moves toward wider use.

7. The weak-link opportunity

AI will not improve every part of every job equally. It may accelerate drafting, analysis, coding, or synthesis while leaving judgment, trust, coordination, or verification as the limiting factor. Strengthening one activity does not automatically improve the larger system.

That is why teams should look at work as connected activities. Where does work break down? Where does information fail to move? Where is expertise unavailable? Where do delays, rework, or noncompliance recur? AI may strengthen one link, reveal that another link is the real constraint, or make it possible to redesign the chain.

Three recurring forms of improvement are especially useful for understanding the model.

Whole-organization improvement

A whole-organization initiative targets a cross-domain outcome, such as reducing rework. Leaders identify the strategic goal or risk. Managers translate it into local context. Individuals and Teams across the organization prototype and validate responses. Evidence moves back toward portfolio decisions, and the necessary expertise, authority, and resources join whatever is selected for broader responsibility.

Domain improvement

A domain initiative targets performance inside a domain, such as improving cycle time, quality, throughput, or resource use. Because the people doing the work may already possess the relevant expertise and authority, they can often prototype, validate, and verify more locally. Additional help becomes necessary when the change affects another domain or requires wider resources and decisions.

Domain governance

A domain-governance initiative targets conformance of AI-generated work in or adjacent to a domain, such as reducing noncompliance. Relevant domain experts create and improve requirements, controls, tests, tools, and escalation paths. Those capabilities improve the quality and capacity of future work rather than operating only as blocking reviews.

These are not exclusive loops. A local improvement may reveal an organization-wide opportunity. A prototype may expose the need for a reusable control. A control may make a new class of experiments practical. The model is valuable when it helps the organization see these connections and learn from them.

8. AI as a new kind of teammate

It is tempting to describe AI as a tool. In many ways, it is. For teams, it can also be useful to think of AI as a new kind of teammate—as long as the metaphor does not imply judgment, responsibility, or authority that AI does not possess.

AI can contribute language, analysis, synthesis, design, code, scenarios, questions, and alternatives. It can help a team see more options, challenge assumptions, and move faster from an idea to something testable.

It can also expand interdisciplinary exploration. A technical team may not have a marketer available for an early idea, and a marketing team may not have an engineer available for every possibility. AI can help either team create a rough representation of unfamiliar work, ask better questions, and make a possibility concrete before the organization commits broader resources.

That representation is not verified expertise. AI needs direction and context, and its output needs evaluation. It can be helpful in one moment and misleading in the next. It can increase the range and speed of exploration while also increasing noise and false confidence.

Teams therefore need a shared process. They should know whether they are clarifying, ideating, developing, or implementing; what they are asking AI to contribute; what evidence is needed; and who must exercise judgment. AI can make teams faster. Metacognition and domain expertise help them decide whether they are moving in a worthwhile direction.

9. What thriving looks like

When organizations get the people side of AI right, the change is visible in the work—not only in an AI usage dashboard.

Individuals and Teams see themselves as participants in improving work. They discover meaningful problems, make ideas tangible, test usefulness with affected people, and communicate what they learned. They understand that a prototype is an invitation to learn, not a claim that a solution is ready for everyone.

Managers become translators, framers, amplifiers, and noise reducers. They help teams connect direction to local reality and help evidence travel without taking ownership of it or acting as kingmakers.

Leaders remain strategic. They read the ecosystem, identify goals and risks, establish boundaries, resource learning, and decide when evidence warrants broader responsibility. They do not confuse visible AI activity with organizational value.

Domain expertise becomes easier to bring into the work because it is not confined to a separate actor. People apply their expertise to discovery, process improvement, requirements, MVP development, verification, and reusable controls as the scope requires.

Governance becomes more enabling. Repeated concerns are translated into approved tools, clearer boundaries, useful requirements, and controls that make correct practice easier. Verification remains proportionate to the risk rather than becoming either ceremonial or universally blocking.

The organization becomes more adaptive. Wider pilots produce renewed learning. Integrated improvements have named ownership, support, monitoring, and a path to retirement. People put AI to work for them, while the organization gains more local discovery without allowing every experiment to become official.

10. Recommendations for organizations

Organizations that want AI to help people thrive should focus on seven priorities.

1. Keep leadership strategic

Monitor the ecosystem, identify organizational goals and risks, make boundaries visible, resource learning, and decide when evidence warrants broader investment. Do not substitute generic usage targets for strategic intent.

2. Measure improved work, not generic AI adoption

Ask what improved, for whom, what was learned, what burdens shifted, and what evidence exists. Use activity measures to diagnose participation, not to claim value.

3. Train for AI-enabled human problem solving

Teach creative problem solving, metacognition, challenge framing, prototyping, validation, evidence communication, and recognition of expertise boundaries. Tailor practice to the activities Leaders, Managers, and Individuals/Teams perform.

4. Use prototypes to learn with affected people

Expect more than idea lists. Let Individuals and Teams own the experiment and its evidence. Invite the people whose work may change to shape validation. Let Managers strengthen framing without taking ownership.

5. Make boundaries visible and use scope to determine collaboration

Clarify prohibited conditions, approved environments, and signals that call for other expertise or authority. Add Individuals or Teams, Manager help, and Leader decisions when affected domains, coordination, resources, or risk expand.

6. Treat governance as enabling infrastructure

Translate strategic risks and domain requirements into tools, controls, tests, and pathways that improve the quality and capacity of future work. Do not make every governance activity a universal gate.

7. Distinguish prototypes, MVPs, and integrated improvements

A prototype is built to learn. An MVP is built and verified for broader responsibility. Wider use requires renewed validation, named ownership, monitoring, support, and a deliberate transition into operating work.

Conclusion: AI value depends on human thinking

AI will continue to improve. Organizations cannot simply wait for it to become powerful enough to solve their problems automatically. The organizations that benefit most will build the human capability to use it well.

AI is not the solution to productivity. It is a new resource for problem solving.

Organizations that understand this distinction will not treat AI adoption as a software rollout or use AI to make the organization about AI. They will use it to help people understand, improve, and reinvent work. People will put AI to work for them.

The future of AI in organizations will not be determined by the technology alone. It will also be determined by the quality of the thinking, evidence, and collaboration people bring to it.